

A Theory of Programming: Denotational, Algebraic and Operational Semantics

C.A.R. Hoare

November 12, 1999

Abstract

Professional practice in a mature engineering discipline is based on relevant scientific theories, usually expressed in the language of mathematics. A mathematical theory of programming aims to provide a similar basis for specification, design and implementation of computer programs. The theory can be presented in a variety of styles, including

1. Denotational, relating a program to a specification of its observable properties and behaviour.
2. Algebraic, providing equations and inequations for comparison, transformation and optimisation of designs and programs.
3. Operational, describing individual steps of a possible mechanical implementation.

This paper presents a simple theory of sequential non-deterministic programming in each of these three styles; by deriving each presentation from its predecessor, mutual consistency is assured.

1 Introduction

A scientific theory is formalised as a mathematical description of some selected class of processes in the physical world. Observable properties and behaviour of such a process can then be predicted from the theory by mathematical deduction or calculation. An engineer applies the theory in the reverse direction. A specification describes the observable properties and behaviour of some system that does not yet exist in the physical world; and the goal is to design and implement a product which can be predicted by the theory to meet the specification.

This paper proposes a mathematical treatment of computer programming in the simple non-deterministic programming language introduced by Dijkstra [?]. The theory is well suited for use by engineers, since it supports both stepwise development of designs from specifications and hierarchical decomposition of complex systems into simpler components which can be designed separately. Furthermore, it permits derivation of a complete set of algebraic laws to help in transformation of designs and optimisation of programs. Finally, an operational semantics is derived; this treats practical aspects of implementation and efficiency of execution.

The insights described here were obtained by a study of communication and concurrency in parallel processes, where the three semantic styles have been applied individually by independent schools of

research to the same class of phenomena. The operational style was used first [?] to define the Calculus of Concurrent Systems (CCS); the algebraic style took precedence in the definition [?] of the Algebra of Concurrent Processes (ACP), whereas the denotational style lies at the basis of the mathematical theory [?] of Communicating Sequential Processes (CSP). Many of the detailed differences between these three process theories originate from their different styles of presentation. To obtain a synthesis based on a full understanding, it is helpful to concentrate on a single theory, and present it fully in all three styles; there is the additional hope that their complementary benefits can be exploited in practice. It is the goal of this paper to explore the relevant techniques in the case of a simple sequential programming language, thereby avoiding any controversy that surrounds the treatment of process algebra.

Not a single idea in this paper is original. The concept of denotational semantics is due to Strachey and Scott [?], and the particular choice of ordering of non-deterministic programs is due to Smyth [?]. The embedding of programs as predicates is due to Hehner [?]. The language is essentially the same as that of Dijkstra [?]. The denotational theory is taken from Tarski's calculus of relations [?]. The treatment of recursion in specifications is given by Tarski's fixed point theorem [?] and for programs by Plotkin [?]. The algebraic treatment of the language has already

been fully covered in [?]. Even the idea of consistent and complementary definitions of programming languages goes back at least to [?].

The only originality in the paper is to show simple ways in which the three presentations of the same language can be derived from each other by mathematical definition, calculation and proof. The denotational theory consists just of a number of separate mathematical definitions of the operators of the language in terms of the second-order predicate calculus. These can be individually formulated and understood in isolation from each other. The algebraic laws can then be derived one by one, without danger of complex or unexpected interactions. A normal form theorem gives insight into the degree of completeness of the laws, and permits additional laws to be proved without induction.

An operational theory is equally easily derived from the algebraic. First an algebraic definition is given for the basic step (transition relation) of an abstract implementation; and then the individual transition rules can be proved separately and individually as algebraic theorems, again with reduced risk of complex or unexpected interactions. The phenomena of deadlock (no transitions) and divergence (an infinite sequence of transitions) are analysed, and shown to relate correctly to their algebraic interpretation.

As always in such smooth developments, the simplicity is an artefact of many laborious and less suc-

cessful iterations, mercifully concealed from the reader. Another reason for the simplicity and modularity of the proofs described above is that they follow the natural progression from abstract description to concrete implementation. It is possible (and indeed more usual) to work in the other direction, starting with an operational presentation. A concept of bisimulation is then selected, permitting the proof of algebraic laws; and a model can then be derived by a standard initial algebra construction. A derivation in both directions establishes completeness as well as consistency of the three presentations. But that is the subject of another paper.

2 Observations and Predicates

When a physical system is described by a mathematical formula, the free variables of the formula are understood to denote results of possible measurements of selected parameters of the system. For example, in the description of a mechanical assembly, it may be understood that x denotes the projection of a particular joint along the x -axis, $\dot{x}$ stands for the rate of change of x , and t denotes the time at which the measurement is taken. A particular observation can be described by giving measured values to each of these variables, for example:

$$x = 14\text{mm} \wedge \dot{x} = 7\text{mm/s} \wedge t = 1.5\text{sec.}$$

The objective of science is not to construct a list of actual observations of a particular system, but rather to describe all possible observations of all possible systems of a certain class. The required generality is embodied in mathematical equations or inequations, which will be true whenever their free variables are given values obtained by particular measurements of any particular system of that class. For example, the differential equation

$$\dot{x} = 0.5 \times x, \qquad \text{for } t \leq 3$$

describes the first three seconds of movement of a point whose velocity varies in proportion to its distance along the x axis. The equation is clearly satisfied by the observation given previously, because

$$7 = 0.5 \times 14 \quad \text{and} \quad 1.5 \leq 3.$$

In applying this insight to computer programming, we shall confine attention to programs in a high level language, which operate on a fixed collection of distinct global variables

$$x, y, \dots z.$$

The values of these variables are observed either before the program starts or after it has terminated. To name the final values of the variables (observed after the program terminates), we place a dash on

the names of the variables

$$x', y', \dots, z'.$$

But to name the initial values of the variables (observed before the program starts), we use the variable names themselves, without decoration. So an observation of a particular run of a program might be described as a conjunction

$$x = 4 \quad \wedge \quad x' = 5 \quad \wedge \quad y' = y = 7.$$

This is just one of the possible observations of a program that adds one to the variable x , and leaves unchanged the values of y and all the other variables; or in familiar symbols, the single assignment

$$x := x + 1.$$

A general formula describing all possible observations of every execution of the above program is

$$x' = x + 1 \quad \wedge \quad y' = y \quad \wedge \dots \wedge \quad z' = z.$$

Such a formula will henceforth be abbreviated by the programming notation which it exactly describes; for example, the meaning of an assignment is actually explained by the *definition*

$$x := x + 1 \quad =_{df} \quad x' = x + 1 \wedge y' = y \wedge \dots \wedge z' = z.$$

Similarly, a program which makes no change to anything is written as II (pronounced “skip”) and defined

$$\text{II} =_{df} x' = x \wedge y' = y \wedge \dots \wedge z' = z.$$

In words, an observation of the final state of II is the same as that of its initial state.

Of course, high level programs are more usually (and more usefully) regarded as instructions to a computer, “*given* certain values of $x, y, \dots, z$, *to find* values of $x', y', \dots, z'$ that will make the predicate true”. But for the purpose of our mathematical theory, there is no need to distinguish between descriptive and imperative uses of the same predicate.

In engineering practice, a project usually begins with a specification, perhaps embodied in a formal or informal contract between a customer and an implementor. A specification too is a predicate, describing the desired (or at least permitted) properties of a product that does not yet exist. For example, the predicate

$$x' > x \quad \wedge \quad y' = y$$

specifies that the value of x is to be increased, and the value of y is to remain the same. No restriction is placed on changes to any other variable. There are many programs that satisfy this specification, including the previously quoted example

$$x := x + 1.$$

Correctness of a program means that every possible observation of any run of the program will yield values which make the specification true; for example, the specification $(x' > x \wedge y' = y)$ is satisfied by the observation $(x = 4 \wedge x' = 5 \wedge y' = y = 7)$. The formal way of defining satisfaction is that the specification is implied by a description of the observation, for example

$$(x = 4 \wedge x' = 5 \wedge y' = y = 7) \Rightarrow (x' > x \wedge y' = y).$$

This implication is true for all values of the observable variables $x, x', y, y', \dots, z, z'$:

$$\forall x, \dots, z' :: (x = 4 \wedge x' = 5 \wedge y' = y = 7) \Rightarrow (x' > x \wedge y' = y).$$

In future, we will abbreviate such universal quantification by Dijkstra's conventional square brackets, which surround the universally qualified formula thus

$$[(x = 4 \wedge x' = 5 \wedge y = y' = 7) \Rightarrow (x' > x \wedge y' = y)].$$

In fact, the specification is satisfied not just by this single observation but by every possible observation of every possible run of the program $x := x + 1$:

$$[x := x + 1 \Rightarrow x' > x \wedge y' = y].$$

This mixture of programming with mathematical no-

tations may seem unfamiliar; it is justified by the identification of each program with the predicate which describes exactly its range of possible behaviours. Both programs and specifications are predicates over the same set of free variables; and that is why the concept of program correctness can be so simply explained as universally quantified logical implication between a program and its specification.

Logical implication is equally interesting as a relation between two programs or between two specifications. If S and T are specifications,

$$[S \Rightarrow T]$$

means that T is a more general or abstract specification than S , and at least as easy to implement. Indeed, by transitivity of implication, any program that correctly implements S will serve as an implementation of T , though not necessarily the other way round. So a logically weaker specification is easier to implement, and the easiest of all is the predicate *true*, which can be implemented by anything.

Similarly, if P and Q are programs,

$$[P \Rightarrow Q]$$

means that P is a more specific or determinate program than Q , and it is (in general) more useful. Indeed, by transitivity of implication, any specification met by Q will be met by P , though not necessarily the other way round. So a logically weaker program

is for any given purpose less likely to serve; and the weakest program **true** is the most useless of all.

The initial specification of a complex product is usually separated from its eventual implementation by one or more stages of development. The interface between each stage can in principle be formalised as a design document D . If this is also interpreted as a predicate, the correctness of the design is assured by the implication

$$[D \Rightarrow S]$$

and the correctness of the later implementation P by

$$[P \Rightarrow D].$$

The correctness of P with respect to S (and the validity of the whole idea of stepwise development) follows simply by transitivity of implication:

$$\text{If } [P \Rightarrow D] \text{ and } [D \Rightarrow S] \text{ then } [P \Rightarrow S].$$

When a predicate is used as a specification, there is no reason to restrict the mathematical notations available for its expression. Indeed, any notation with a clear meaning should be available, because clarity of specification is the only protection we have against subsequent misunderstandings of the client's requirements, which can often lead to disappointment or even rejection of a delivered product.

Particularly important aids to clarity of specification are the simple connectives of Boolean algebra, conjunction (and), disjunction (or), and negation (not). Conjunction is needed to connect individual requirements such as “Temperature must be less than 30° *and* more than 27°”. Disjunction is needed to provide useful options for economic implementation: “For mixing, use either the pressure vessel *or* the settling tank”. And negation is needed for even more important reasons: “It must *not* explode”.

The freedom of notation which is appropriate for specification cannot be extended to the programming language in which the ultimate implementation is expressed. Programming notations must be selected to ensure computability, compilability, and reasonable efficiency of execution. In a given programming language, there is a limited collection of combinators available for construction of programs from their primitive components. Typical components include assignments, inputs and outputs; and typical combinations include conditionals, sequential composition, and some form of iteration or recursion. It is for good reason that most programming languages exclude the Boolean combinators and quantifiers of mathematical logic. For example, there is no programming language or compiler that would enable you to protect against disaster by writing a program that *causes* an explosion and then avoid explosion by just negating the program before execution.

A result of these practical restrictions is that, although we can interpret all programs as predicates, the converse is obviously invalid: not every predicate describes the behaviour of a program. For example, consider the extreme predicate *false*. No observation satisfies this predicate, so the only object that it could correctly describe is one that gives rise to no observation whatsoever. From a scientific viewpoint, such an object does not exist, and could never be constructed. The notations of a programming language must therefore be defined to ensure that they can never express the predicate *false*, or any other wholly unimplementable predicate.

This means that we must live with the danger of proposing and accepting an unimplementable predicate as a specification. Indeed, any general notational restriction that ensures computability (or even just satisfiability) could seriously impact clarity and conciseness of specification, and so increase the much greater risk of failure to capture the true requirements and goals of the project. Once these have been correctly formalised, a check on implementability, and on efficiency of implementation, may be made separately with the aid of mathematics or good engineering judgement; and this will be confirmed in the end by successful delivery of an actual product which meets the specification. There is fortunately no danger whatsoever of delivering an implementation of an unimplementable specification.

3 The programming language

In this section we shall give a denotational semantics of our simple sequential programming language in terms of predicates describing the behaviour of any program expressed in that language. As explained earlier, the variables $x, y, \dots, z$ stand for the initial values of the like-named global variables of the program, and $x', y' \dots, z'$ stand for the final values.

Let $e, f, \dots, g$ stand for expressions such as $x + 1, 3 \times y + z, \dots$ that can feature on the right hand side of an assignment. Clearly, their free variables are confined to the undashed variables of the program; and for simplicity, we assume that all expressions always evaluate successfully to a determinate result. Generalising an example given earlier, we define a simple assignment,

$$x := e \quad =_{df} \quad x' = e \wedge y' = y \wedge \dots \wedge z' = z.$$

The program which makes no change is just a special case

$$\text{II} \quad =_{df} \quad x := x.$$

A multiple assignment has a list of variables on the left hand side, and a list of the same number of expressions on the right; it is defined

$$x, y := e, f \quad =_{df} \quad x' = e \wedge y' = f \wedge \dots \wedge z' = z.$$

A clear consequence of the definition is that an implementation must evaluate all the expressions on the right hand side before assigning any of the resulting values to a variable on the left hand side.

Other consequences can be simply formulated as algebraic laws; they have very simple proofs. For example

$$\begin{aligned}x &:= e &= x, y &:= e, y \\x, y &:= e, f &= y, x &:= f, e.\end{aligned}$$

All the definitions and laws extend to lists of more than two variables, for example

$$(z, y := g, f) = (x, y, \dots, z := x, f, \dots, g).$$

In fact every assignment may be transformed by these laws to a *total* assignment

$$x, y, \dots, z := e, f, \dots, g$$

where the left hand side is a list of all the free variables of the program, in some standard order. In future we will abbreviate this to

$$v := f(v)$$

where v is the vector $(x, y, \dots, z)$ of program variables, and f is a total function from vectors to vectors. Predicates will be similarly abbreviated

$$P(v, v') \text{ instead of } P(x, y, \dots, z, x', y', \dots, z').$$

Any non-trivial program is composed from its primitive components by the combining notations (combinators) of the programming language. The runtime behaviour of a composite program is obtained by actual execution of its components — all, some, or sometimes even none of them. Consequently, at a more abstract level, a predicate describing this composite behaviour can be defined by an appropriate composition of predicates describing the individual behaviours of the components. So a combinator on programs is defined as a combinator on the corresponding predicates.

The first combinator we consider is the *conditional*. Let b be a program expression, containing only undashed variables and always producing a Boolean result (true or false); and let P and Q be predicates describing two fragments of program. A conditional with these parameters describes a program which behaves like P if b is initially true, and like Q if b is initially false. It may therefore be defined

$$P \triangleleft b \triangleright Q \quad =_{df} \quad (b \wedge P) \vee (\neg b \wedge Q).$$

A more usual notation for a conditional is

$$\mathbf{if } b \mathbf{ then } P \mathbf{ else } Q \text{ instead of } P \triangleleft b \triangleright Q.$$

The reason for the change to infix notation is that it simplifies the expression of algebraic laws:

$$P \triangleleft b \triangleright P = P$$

$$P \triangleleft b \triangleright Q = Q \triangleleft \neg b \triangleright P$$

$$\begin{aligned} (P \triangleleft b \triangleright Q) \triangleleft b \triangleright R &= P \triangleleft b \triangleright (Q \triangleleft b \triangleright R) \\ &= P \triangleleft b \triangleright R \end{aligned}$$

$$P \triangleleft b \triangleright (Q \triangleleft c \triangleright R) = (P \triangleleft b \triangleright Q) \triangleleft c \triangleright (P \triangleleft b \triangleright R).$$

The first law expresses idempotence, the second gives a form of skew symmetry, the third is an associative law, and the fourth states the distribution of any conditional operator $\triangleleft b \triangleright$ through the conditional $\triangleleft c \triangleright$, for any condition c . All the laws may be proved by propositional calculus; the easiest way is to consider separately the cases when b is true and when it is false. In the first case, replace $P \triangleleft b \triangleright Q$ by P and in the second case by Q . The purpose of the algebraic laws is to help in mathematical reasoning, without such tedious case analyses.

The most characteristic combinator of a sequential programming language is sequential composition, often denoted by semicolon. $(P ; Q)$ may be executed by first executing P and then Q . Its initial state is that of P , and its final state is that of Q . The final state of P passed on as the initial state of Q ; but this is only an intermediate state of $(P ; Q)$, and it cannot be directly observed. All we know is that it exists. The formal definition therefore uses existential quantification to hide the intermediate observation, and to remove the variables which record it from the list

of free variables of the predicate.

$$P(v, v') ; Q(v, v') =_{df} \exists v^0 P(v, v^0) \wedge Q(v^0, v').$$

Here, the vector variable v^0 stands for the correspondingly decorated list of bound variables

$$(x^0, y^0, \dots, z^0).$$

These record the intermediate values of the program variables

$$(x, y, \dots, z),$$

and so represent the intermediate state as control passes between P and Q . But this operational explanation is far more detailed than necessary. A clever implementation is allowed to achieve the defined effect by more direct means, without ever passing through any of the possible intermediate states. That is the whole purpose of a more abstract definition of the programming language.

In spite of the complexity of its definition, sequential composition obeys some simple, familiar and obvious algebraic laws. For example, it is associative and has II as its left and right unit. Finally, sequential composition distributes leftward (but not rightward) over the conditional. This asymmetry arises because the condition b is allowed to mention only the initial values of the variables, and not the final (dashed) variables.

$$\begin{aligned}
(P ; Q) ; R &= P ; (Q ; R) \\
\Pi ; P &= P = P ; \Pi \\
(P \triangleleft b \triangleright Q) ; R &= (P ; R) \triangleleft b \triangleright (Q ; R).
\end{aligned}$$

If e is any expression (only mentioning undashed variables), the assignment

$$x := e$$

changes the value of x so that its final value is the same as the initial value of e , obtained by evaluating e with all its variables taking its initial values. So if $P(x)$ is any predicate mentioning x , P is true of the final value of x in just the case that P is true of e , i.e.,

$$\begin{aligned}
x := e ; P(x) &= (\exists x^0 : x^0 = e : P(x^0)) \\
&= P(e).
\end{aligned}$$

But $P(e)$ is just P with x substituted by e . This substitution effect generalises to any expression:

$$(x := e ; f(x)) = f(e).$$

For example

$$(x := x + 1 ; (3 \times x + y < z)) = (3 \times (x + 1) + y < z).$$

This convention permits a rightward distribution law for conditionals:

$$x := e ; (P \triangleleft b \triangleright Q) = (x := e ; P) \triangleleft x := e ; b \triangleright (x := b ; Q).$$

Let P and Q be predicates describing the behaviour of programs. Their disjunction $(P \vee Q)$ describes the behaviour of a program which may behave like P or like Q , but does not say which it will be. As an operator of our programming language, disjunction may be easily implemented by arbitrary selection of either of the operands; and the selection may be made at any time, either before or after the program is compiled or even after it starts execution. Disjunction is an extremely simple explanation of the traditionally obscure phenomenon of non-determinism in computing science; and its simplicity provides additional justification for the definition and manipulation of programs as predicates.

All the program combinators defined so far distribute through disjunction. This means that separate consideration of each case is adequate for all reasoning about non-determinism. Curiously, disjunction also distributes through itself and through the conditional

$$\begin{array}{llll}
P \triangleleft b \triangleright (Q \vee R) & = & (P \triangleleft b \triangleright Q) & \vee & (P \triangleleft b \triangleright R) \\
P ; (Q \vee R) & = & (P ; Q) & \vee & (P ; R) \\
(Q \vee R) ; P & = & (Q ; P) & \vee & (R ; P) \\
P \vee (Q \vee R) & = & (P \vee Q) & \vee & (P \vee R) \\
P \vee (Q \triangleleft b \triangleright R) & = & (P \vee Q) & \triangleleft b \triangleright & (P \vee R).
\end{array}$$

As a consequence of distribution through disjunction, all program combinators also share the property of monotonicity. A function f is said to be *mono-*

tonic if it preserves the relevant ordering, in this case implication. More formally

$$[f.X \Rightarrow f.Y] \text{ whenever } [X \Rightarrow Y].$$

(Here, X and Y are mathematical variables ranging over *predicates*, and the line displayed above is true, no matter what predicates take the place of X and Y). All program combinators defined so far are monotonic in all arguments; for example

$$[X ; Y \Rightarrow X' ; Y'] \text{ whenever } [X \Rightarrow X'] \text{ and } [Y \Rightarrow Y'].$$

Monotonicity is a very important principle in engineering. Consider an assembly which tolerates a given range of variation in its working environment. Consider also some one of its components, which also has a certain tolerance t . The tolerance of the whole assembly can be expressed as some function f of t . The engineer usually assumes that f is a monotonic function, so that if the component is replaced by one

with a broader tolerance t' , then the tolerance of the whole assembly will in general also be broader, or at worst, the same:

$$[t \leq t' \Rightarrow f(t) \leq f(t')].$$

Problems arising from violation of monotonicity are in practice the most difficult to diagnose and rectify, because they invalidate the whole theory upon which design of the assembly has been based.

When faced with the task of implementing a complex specification S , it is usual to make an early decision on the general structure of the product, for example as the sequential composition of two program components. To formalise and communicate this decision, each of these components is going to need separate specifications, say D and E . The correctness of these specifications can be checked before implementation by proof of the implication

$$[(D ; E) \Rightarrow S], \quad (*)$$

where the sequential composition between specifications has the same definition as between programs considered as predicates. Now what remains is the presumably simpler task of finding two programs P and Q which implement the two designs, i.e.,

$$[P \Rightarrow D] \text{ and } [Q \Rightarrow E].$$

Now all that remains is to deliver the product $(P ; Q)$.

By monotonicity of sequential composition

$$[P ; Q \Rightarrow D ; E],$$

and the fact that

$$[(P ; Q) \Rightarrow S]$$

follows by transitivity from a proof of the correctness of the design step (*). What is more, this proof was completed before the start of implementation of P or Q . The technique can be repeated on the components P and Q ; and because of monotonicity it extends to all other program combinators. Their monotonicity is essential to the general engineering method of stepwise design decomposition. Note that designs are expressed in a mixture of programming notations (for decisions that have already been taken) and more general predicates (for the parts that are specified but still need to be designed). This is yet another advantage of the philosophy of expressing both programs and specifications in the same logical space of predicates.

4 Recursion

A final advantage of monotonicity is that it permits a simple treatment of the important programming concept of recursion and of its important special case, iteration; without this, no program can take longer to execute than to input. Predicates over a given set

of observational variables may be regarded as a complete lattice under implication ordering, with universal quantification as meet and existential as join. The bottom of the lattice is the strongest predicate *false* and the top is **true**. Here we will use bold font to distinguish **true** (considered as a program predicate over free variables v, v') from italic *true*, which is a possible value of a Boolean expression b (containing only free variables v).

Moving to a second-order predicate calculus, we introduce a variable X to stand for an arbitrary predicate over the standard set of first-order variables. Fortunately, all the combinators of our programming language are monotonic, and any formula constructed by monotonic functions is monotonic in all its free variables. Let $G.X$ be a predicate constructed solely by monotonic operators and containing X as its only free predicate variable. Tarski's theorem [?] guarantees that the equation

$$X = G.X$$

has a solution for X ; and this is called a fixed point of the function G . Indeed, among all the fixed points, there is a weakest one in the implication ordering. This will be denoted by

$$(\mu X :: G.X).$$

It can be implemented as a single non-recursive call of a parameterless procedure with name X and body

$(G.X)$. Occurrences of X within $(G.X)$ are implemented as recursive calls on the same procedure.

The mathematical definition of recursion is given by Tarski's construction:

$$\mu X :: G.X =_{df} \vee \{X : [X \Rightarrow G.X] : X\}$$

where $\vee$ is the lattice join applied to the set of all solutions of $(X \Rightarrow G.X)$. The following laws state that the join is indeed a fixed point of G , and that it is the weakest such.

$$\begin{aligned} [G.(\mu X :: G.X) &\equiv (\mu X :: G.X)] \\ [Y \Rightarrow \mu X :: G.X] &\text{ whenever } [Y \Rightarrow G.Y]. \end{aligned}$$

A simple common case of recursion is the iteration or while loop. If b is a condition,

while b **do** P

repeats the program P for as long as b is true before each iteration. More formally, it can be defined as the recursion

$$(\mu X :: (P ; X) \triangleleft b \triangleright \text{II}).$$

An even simpler example (but hopefully less common) is the infinite recursion which never terminates

$$\mu X.X.$$

This is the weakest solution of the trivial equation

$$X = X$$

and is therefore the weakest of all predicates, namely **true**. In engineering practice, a non-terminating program is the worst of all programs, and must be carefully avoided by any responsible engineer. That will have to suffice as justification for practical use of a theory which equates any non-terminating program with a totally unpredictable one, which is the weakest in the lattice ordering.

Consider now the program

$$(\mu X :: X) ; x, y, \dots, z := 3, 12, \dots, 17$$

which starts with an infinite loop. In any normal implementation, this would fail to terminate, and so be equal to $(\mu X :: X)$. Unfortunately, our theory gives the unexpected result

$$x' = 3 \ \wedge \ y' = 12 \ \wedge \ \dots \wedge \ z' = 17,$$

the same as if the prior non-terminating program had been omitted. To achieve this result, an implementation would have to execute the program backwards, starting with the assignment, and stopping as soon as the values of the variables are known. While backward execution is not impossible (indeed, it is standard for lazy functional languages), it is certainly not efficient for normal procedural languages. Since we

want to allow the conventional forward execution, we are forced to accept the practical consequence that the program

$$(\mu X :: X); P$$

will fail to terminate for any program P ; and the same is true of

$$P; (\mu X :: X).$$

Substituting $(\mu X :: X)$ by its value **true** we observe in practice of all programs P that

$$\begin{aligned} \mathbf{true}; P &= \mathbf{true} \\ P; \mathbf{true} &= \mathbf{true}. \end{aligned}$$

These laws state that **true** is a zero for sequential composition.

But these laws are certainly not valid for an arbitrary predicate P . As always in science, if a theory makes an incorrect prediction of the behaviour of an actual system, it is the theory that must be adapted; and this usually involves an increase in complication. That is what requires and justifies introduction of new concepts and variables, which cannot perhaps be directly observed or controlled, but which are needed to explain what would otherwise be anomalies in more directly observable quantities. All the discoveries of fundamental forces and particles in modern physics have been made in this way.

In the case of computer programs, the anomaly is resolved by investigating more closely the phenomena of starting and stopping of programs. The collection of free variables describing programs is enlarged to include two new Boolean variables, which are never allowed to appear in the text of the program:

st , which becomes true when the program has been started, and is false beforehand.

st' , which becomes true when the program has stopped, and remains forever false in case of non-termination (and *a fortiori*, if program is never started).

While st' is false, the final values of the program variables are unobservable, and the predicate describing the program should make no prediction about these values. Similarly, while st is false, even the initial values are unobservable. These considerations underlie the validity of the desired zero laws.

We still maintain the convention that no observation will be made of the variables while the program is running, so we never observe that st is true and st' is false, except in the case of non-termination. This is the essential abstraction from details of execution time, which permits a separation of concerns between correctness and efficiency in reasoning about program behaviour. It also permits programs written for the IBM 704 in 1960 to run correctly on supercomputers

for the present day, in spite of a vast difference in speed.

The variables st and st' are useful also in specifications of components of larger programs. The correctness and even the termination of a component with specification Q is often dependent on some assumed properties of the initial values of the variables. This assumption is described by a *precondition* P , which will be true before the program starts. The specification can then be written

$$(st \wedge P) \Rightarrow (st' \wedge Q)$$

or in words “If the program components start in a state satisfying P , it will stop in a state satisfying Q .”

The responsibility for ensuring that P is true at the start is thereby delegated to the preceding part of the program. If the assumption is violated, no constraint whatsoever is placed on the designed behaviour of the subsequent program; it may even fail to terminate. Successful teamwork in a large engineering project always depends on appropriately selected assumptions made by the individual designers, and the corresponding obligations undertaken by their colleagues. So it is worth while to introduce a special notation

$$(P, Q) =_{df} (st \wedge P \Rightarrow st' \wedge Q).$$

This is the primitive notation used by Morgan in [?].

The clear distinction of precondition P from postcondition Q is a distinctive feature of VDM [?].

Another advantage of explicit mention of starting and stopping is a solution of the postponed problem of undefined expressions in assignments. For each expression e of a reasonable programming language, it is possible to calculate a condition $\mathcal{D}e$ which is *true* in just those circumstances in which e can be successfully evaluated. For example

$$\begin{aligned}\mathcal{D}17 &= \mathcal{D}x = \textit{true} \\ \mathcal{D}(e + f) &= \mathcal{D}e \wedge \mathcal{D}f \\ \mathcal{D}(e/f) &= \mathcal{D}e \wedge \mathcal{D}f \wedge (f \neq 0).\end{aligned}$$

Successful execution of an assignment relies on the assumption that the expression will be successfully evaluated, so we formulate a *new* definition of assignment

$$x := e \quad =_{df} \quad (\mathcal{D}e \wedge x' = e \wedge y' = y \wedge \dots \wedge z' = z).$$

Expressed in words, this definition states that

- either the program has not started ($st = false$)
and nothing can be said about its initial
and final values
- or the initial values of the variables are such
that evaluation of e fails ($\neg \mathcal{D}e$), and noth-
ing can be said about the final values
- or the program has terminated ($st' = true$),
and the value of x' is e , and the final values
of all the other variables are the same as
their initial values.

Fortunately, this is the only new definition that is needed; the definition of conditionals, recursion, and sequential composition remain unchanged, and all laws (except those involving assignment) remain valid. In fact, the laws involving assignment also remain valid, provided that their variables range not over arbitrary predicates, but only over predicates expressed in programming notations. For this restricted class of predicates (hereafter called programs), we will have to prove the unit laws

$$\text{II}; P = P = P; \text{II}, \text{ for all programs } P$$

as well as the new zero laws

$$P; \mathbf{true} = \mathbf{true} = \mathbf{true}; P, \text{ for all programs } P.$$

In compensation, the zero laws give an assurance that

no programs can be equal to the unimplementable predicate *false*, which does not satisfy them.

It is quite easy to check that the zero and unit laws are valid for the simple case of programs that are assignments, even when these are interpreted according to the new definition. This proof can be extended by structural induction to more complex kinds of program. A simple example of a lemma needed in this proof is:

If P and Q satisfy the laws

$$P; \mathbf{true} = \mathbf{true} = Q; \mathbf{true}$$

then so do $(P; Q)$ and $(P \triangleleft b \triangleright Q)$.

The proof of this theorem is also quite simple:

$$\begin{aligned} (P; Q); \mathbf{true} &= P; (Q; \mathbf{true}) = P; \mathbf{true} = \mathbf{true} \\ (P \triangleleft b \triangleright Q); \mathbf{true} &= (P; \mathbf{true}) \triangleleft b \triangleright (Q; \mathbf{true}) \\ &= \mathbf{true} \triangleleft b \triangleright \mathbf{true} = \mathbf{true}. \end{aligned}$$

Unfortunately the required additional theorems for the left zero law and for the recursion operator μ are much more difficult to prove. The relevant mathematics is worked out in the next chapter.

References

- [1] E.W. Dijkstra, “A Discipline of Programming”, Prentice Hall, 1976.
- [2] A.R.J.G. Milner, “A Calculus of Communicating Systems”, LNCS 92, Springer-Verlag, 1980.
- [3] J.A. Bergstra and J.W. Klop, “Algebra of Communicating Processes with Abstraction”, Theoretical Computer Science 37(1), 77-121, 1985.
- [4] S.D. Brookes, C.A.R. Hoare and A.W. Roscoe, “A Theory of Communicating Sequential Processes”, Journal of ACM 31(7) 560-599, 1984.
- [5] D.S. Scott and C. Strachey, “Towards a Mathematical Semantics for Computer Languages”, PRG-6, Oxford 1971.
- [6] M.B. Smyth, “Power domains”, JCSS (16) 23-26, 1978.
- [7] E.C.R. Hehner, “Predicative Programming”, Comm ACM 27(2), 134-143.
- [8] A. Tarski, “On the Calculus of Relations”, J Symbolic Logic 6, 73-89, 1941.
- [9] G.D. Plotkin, “A Structural Approach to Operational Semantics”, Report DAIMI-FN-19, Computer Science Department, Aarhus University, 1981.
- [10] A. Tarski, “A Lattice-theoretic Fixed Point Theorem and its Applications”.

- [11] C.A.R. Hoare et al., “The Laws of Programming”, Comm ACM 30(8), 672-87.
- [12] C.A.R. Hoare, R.E.Lauer, “Consistent and complementary formal theories of the semantics of programming languages”, Acta Informatica 3(2), 135-153, 1974.
- [13] C.B. Jones, “Systematic Software Development using VDM”, Prentice Hall International, 1986.
- [14] C.C. Morgan, “Programming from Specifications”, Prentice Hall International, 1990.
- [15] J.A. Goguen and T. Winkler, “Introducing OBJ3”, Technical Report SRI-CSL-88-9, SRI International Computer Science Lab., 1988.